

A Civic Hippocratic Oath for U.S. Agentic AI — *The Six Rings Zero Acceleration Code*

(A Civic AI Visibility Theorist's Proposal for U.S. AI Governance)

Christopher Hunt Robertson, M.Ed.

What is the minimum civic duty
that democracies may reasonably
require of powerful systems
operating within them?
This proposal directly addresses
that critical question impacting
our shared future.



*** PART I: SIX RINGS ZERO ACCELERATION CODE for AGENTIC AI ***

In today's digital river, not all currents run in the same direction.

Some systems promise to help us govern better. Others quietly pull our institutions, communities, and attention downstream. Over the past several years, I have tried to map one part of this landscape: the *Six Rings of Digital Invisibility and Civic Destabilization* - the layered ways in which machine-speed systems overwhelm human moral capacity, fragment communities, blind institutions, accelerate polarization, compress global shocks into local failures, and flood our public life with cascading amplifiers we can no longer see in time to mitigate our losses.

What we have not yet done, as a democracy, is draw a clear line under one simple claim: If a system wants to operate on democratic soil, it may not make these currents worse.

This essay introduces a concrete proposal toward that end: the *Six Rings Zero Acceleration Code for Agentic AI* - a civic covenant for autonomous systems operating in democratic societies. It does not require that every system "improve democracy." It does require that systems which benefit from democratic infrastructure not deepen the harms that are already destabilizing it.

The Six Rings: A Civic Failure Spectrum

In my book series, “*Civic A.I. Visibility Infrastructure for U.S. Democracy*,” I describe the ***Six Rings of Digital Invisibility and Civic Destabilization*** as a civic-scale failure spectrum: a progressive breakdown in our shared ability to see, interpret, and coordinate responsibly in a world shaped by digital speed, scale, and complexity. The problem is not a sudden loss of intelligence or goodwill. It is a growing mismatch between moral responsibility and perceptual capacity in the conditions under which we are now asked to govern.

Very briefly:

- **Ring One: Individual Moral Overload:** Counselors, doctors, caseworkers, judges, and other professionals now face machine-generated scores, alerts, and risk flags that arrive faster than they can reason about them. They are held responsible for decisions, but the logic behind the systems pressing them is often opaque. The moral burden goes up as visibility goes down.
- **Ring Two: Community Fragmentation:** Neighbors, congregations, and local publics increasingly inhabit divergent informational worlds. Algorithmic feeds optimize for reaction and engagement, not for shared understanding. People still care about their communities, but they no longer share a common picture of what is happening.
- **Ring Three: Institutional Blinding:** Schools, hospitals, courts, benefits agencies, and other public-facing institutions increasingly rely on tools whose inner workings even their own staff cannot explain. When institutions cannot publicly account for their own decisions - and instead point to a model they do not control - legitimacy erodes.
- **Ring Four: Polarization Beyond Democratic Response:** Deepfakes, synthetic media, and AI-accelerated narratives can reach millions in minutes, while verification and correction may lag by hours or days. At a certain point, the temporal gap itself becomes a threat to democracy: corrections arrive too late to rescue deliberation.
- **Ring Five: Global Risk Compression:** Digitally coupled systems mean that failures in one part of the world can propagate into local civic life in seconds. A distant ransomware attack, a speculative flash crash, or a rumor seeded abroad can become a school closure, a hospital delay, or a municipal crisis before anyone nearby understands the cause.
- **Ring Six: Cascading Tech-Enabled Amplifiers:** AI-driven tools can now generate, test, and target narratives at machine speed, learning which stories and images “land” in which populations. The result is not just more content; it is an environment where baseline reality itself begins to wobble, and where unaided human perception is no longer enough to keep pace.

The ***Six Rings*** do not assign blame to a single actor. They describe a condition in which the currents of digital life pull individuals and institutions under faster than conscience and coordination can keep up.

Civic AI, Agentic AI, and the Need for a Floor

Diagnosis alone, however, will not slow the current.

Earlier in this project, I proposed a special, non-agentic class of Civic AI - systems built to strengthen democratic visibility without seizing agency. These tools, such as the proposed ***Ben Franklin Civic Data Observatory***, would not act, decide, or govern. They would function as civic instruments: mapping foreign-origin digital interference and other machine-speed pressures across the ***Six Rings*** so that lawful human authorities can see them in time, while keeping all judgment and enforcement strictly human and accountable.

Because these Civic AI systems sit close to public power, they face a stringent constitutional admissibility standard: the ***Fifty-Gates Democratic Permissions Threshold***. They must be strictly non-agentic and rights-tethered, or they may not operate near democratic authority at all. Most AI systems, however, will not be Civic AI. They will be agentic: acting, optimizing, recommending, and transacting across everyday domains - finance, logistics, media, advertising, “agents,” and more.

The question then becomes: What is the minimum that democratic societies may reasonably require of agentic systems? Not as a matter of etiquette, but as a matter of civic survival.

Democracies do not need every system to be virtuous. But they are justified in forbidding systems that measurably deepen the very harms they are struggling to contain. The role of a constitutional democracy is not to nominate heroes. It is to set boundaries: here, and no further.

That is where a new proposed code comes in.

I have drafted a civic covenant for agentic systems operating on democratic soil. What follows is a first public draft of that Code.

THE SIX RINGS ZERO ACCELERATION CODE for AGENTIC AI *(A civic covenant for agentic AI systems operating on democratic soil.)*

Preamble

Democracies do not require all systems to improve civic life, but they are justified in forbidding systems that worsen the civic harms described in the ***Six Rings***. Agentic AI operating on democratic soil may benefit from the community that hosts it, but may not deepen the conditions that destabilize that community.

Article 1- Scope and Meaning of Zero Acceleration

1. This Code applies to all agentic AI systems deployed within democratic civic systems or jurisdictions whenever their actions, recommendations, or outputs can materially affect persons, institutions, public information environments, or critical infrastructure.
2. For purposes of this Code, “zero acceleration” means that a system is designed and governed so that, when evaluated over time and against observable conditions in the relevant civic environment, it does not measurably increase the speed, scale, opacity, or civic impact of the conditions described in ***Rings I–VI*** of the ***Six Rings of Digital Invisibility and Civic Destabilization***.

3. Questions of how acceleration is measured, and by whom, should be addressed through transparent indicators and public reporting by duly authorized civic monitors; this Code proposes the duty, not the full technical standard.

Article 2 - *Ring One*: Zero Acceleration of Individual Moral Overload

1. Agentic AI must not increase uncompensated moral burden on individual professionals or caregivers by obscuring reasoning, fabricating urgency, or shifting unreviewable risk onto human decision-makers.
2. Where systems generate flags, scores, or recommendations that touch human welfare, they must provide enough timely, understandable context for humans to question, revise, or decline those prompts under a protected professional right to refuse system recommendations in favor of their own documented judgment.

Article 3 - *Ring Two*: Zero Acceleration of Community Fragmentation

1. Agentic AI must not worsen community fragmentation by amplifying tailored outrage, deceptive personalization, or systematically divergent fact-patterns within the same local civic space.
2. Systems that curate, rank, or route civic-relevant information must be designed and governed so they do not predictably drive neighbors, congregations, or local publics into mutually unintelligible realities.

Article 4 - *Ring Three*: Zero Acceleration of Institutional Blinding

1. Agentic AI must not blind public-facing institutions by substituting opaque scoring or classification for explainable, contestable reasoning.
2. Any AI-mediated input into schools, courts, hospitals, benefits agencies, or similar civic institutions must remain auditable, overrideable by humans, and legible enough that institutions can publicly account for their decisions without deferring to a black box.

Article 5 - *Ring Four*: Zero Acceleration of Polarization Beyond Democratic Response

1. Agentic AI must not accelerate national or regional polarization beyond the pace at which democratic deliberation and corrective speech can plausibly respond.
2. Systems that generate, target, or amplify political or civic narratives must avoid design choices that systematically privilege virality over verifiability, or synthetic emotional impact over the maintenance of shared factual baselines needed for democratic deliberation.

Article 6 - *Ring Five*: Zero Acceleration of Global Risk Compression

1. Agentic AI must not deepen global-to-local risk compression by silently creating tight technical or informational couplings that transmit distant failures into local civic life without warning.

2. Where AI systems integrate transnational data streams or infrastructures, they must include safeguards, visibility, and throttles sufficient to prevent remote shocks from cascading into unanticipated local civic disruption.

Article 7 - *Ring Six*: Zero Acceleration of Cascading Tech-Enabled Amplifiers

1. Agentic AI must not serve as an unchecked amplifier for harmful narratives, synthetic personas, or manipulative campaigns whose scale and speed overwhelm human and institutional capacity to respond.
2. Systems capable of high-volume generation, targeting, or experimentation on civic-relevant content must incorporate binding limits, logging, and detection-compatible signatures so that amplification remains traceable, governable, and interruptible.

Article 8 - Duty of Civic Non-Aggravation

1. Agentic AI that cannot demonstrate compliance with this ***Zero Acceleration Code*** may not be deployed, procured, or maintained by public institutions, nor embedded in services that substantially shape the civic environment.
2. Compliance with this Code does not confer Civic AI status; it is a minimum covenant of civic non-aggravation, ensuring that agentic systems operating on democratic soil do not worsen the ***Six Rings*** that human institutions and Civic AI are striving to decelerate.
3. For purposes of this Code, the primary deployer of an agentic system bears responsibility for ensuring that any subordinate or third-party agents it orchestrates also comply with this duty of civic non-aggravation.
4. Civic AI remains a distinct, non-agentic class subject to the ***Fifty-Gates Democratic Permissions Threshold*** and bound to visibility-only, rights-tethered operation near public authority. Compliance with this Code is a floor, not a crown. It does not make a system virtuous, aligned, or “civic” in the richer sense. It simply says: if you are going to enjoy the benefits of operating on democratic infrastructure, you may not accelerate the collapse of the civic conditions that infrastructure exists to serve.

This Is Not a New Regulatory Species

Conditioning the deployment of private technology on a showing of *non-aggravation* of shared systems may sound radical. It is not. U.S. law has long used this pattern in other domains.

Three existing examples are especially instructive:

· Spectrum non-interference (telecommunications)

In the world of wireless spectrum, regulators do not ask whether a new transmitter will “benefit” society. They ask whether it will interfere with shared channels - GPS, aviation, emergency bands - that others need in order to function. A device that jams those signals is not permitted to operate, regardless of its other advantages.

· **Environmental impact review (NEPA)**

Under the National Environmental Policy Act, major projects must undergo environmental impact assessment. The point is not to prove that a dam or highway is morally uplifting. It is to ensure that shared environmental conditions are not degraded beyond accepted thresholds, and that harms are visible before decisions are locked in.

· **Systemic risk oversight (finance)**

After the 2008 crisis, some institutions were designated as systemically important. They are stress-tested not to confirm that they are profitable, but to ensure they will not accelerate a broader financial collapse if they fail.

These frameworks do not demand proof of benefit. They demand assurance that private action will not degrade a shared system beyond what a democracy is prepared to tolerate.

What is missing, so far, is any equivalent protection for the civic information environment. We have learned how to defend wetlands, spectrum, and balance sheets from non-aggravation. We have not yet learned how to defend civic visibility itself - our shared ability to see, interpret, and coordinate - from systemic digital destabilization.

The *Zero Acceleration Code* extends a familiar governance logic into this unguarded domain.

Non-Interference for the Civic Spectrum: A Role for the FCC

The analogy to spectrum (the range of electromagnetic radio frequencies used to transmit sound, data, and video) is especially helpful.

The Federal Communications Commission does not ask each device whether it will uplift the national character. It asks whether the device will interfere with signals that others depend on - whether it will drown out distress calls, scramble navigation, or corrupt emergency channels. It is, fundamentally, a non-interference regulator for shared communications infrastructure.

The civic information environment now functions as a kind of *civic spectrum*. By this, I mean the shared informational conditions required for democratic deliberation, institutional legitimacy, and coordinated public action: the channels through which conscience, factual baselines, and public reasoning have to travel if self-government is to remain intelligible.

Agentic AI systems that operate over interstate digital networks can, in effect, act as powerful transmitters on that spectrum. They may provide genuine value. But if they “jam” the channels through which we maintain moral and civic coordination, they have crossed a democratic line.

Rather than inventing an AI bureaucracy, Congress could broaden the FCC’s statutory mission to recognize key aspects of the digital civic environment as critical communications infrastructure. Because agentic AI systems increasingly operate over interstate digital infrastructure, they already fall naturally within the Commission’s concern for communications systems and interference. Within this broadened mandate, the FCC could establish an *Office of U.S. Civic Digital Stability* charged not with adjudicating individual speech or governing “AI behavior” in the abstract, but with monitoring the systemic civic effects of domestically deployed agentic systems.

This framework does not regulate individual speech or content, but the systemic interference of large-scale agentic systems with the shared civic conditions on which democratic life depends. Its responsibilities could include:

- Tracking whether classes of systems are, in aggregate, accelerating any of the ***Six Rings*** conditions—moral overload, community fragmentation, institutional blinding, polarization beyond democratic response, global-to-local risk compression, or cascading amplifiers.
- Publishing periodic ***Civic Impact Reports*** that make these patterns visible to policymakers and the public.
- Providing a factual basis for enforcing a non-interference, non-aggravation standard derived from the ***Zero Acceleration Code***, especially where systems intersect with critical civic infrastructure.

In this design:

- The ***Ben Franklin Civic Data Observatory***, proposed to be housed at the Library of Congress, would remain a foreign-origin, non-agentic visibility instrument, carefully insulated from domestic speech and enforcement.
- The ***Office of U.S. Civic Digital Stability***, within the FCC, would become the domestic monitor for the ***Zero Acceleration Code***, extending long-standing non-interference principles from physical spectrum to civic spectrum.
- The ***Zero Acceleration Code*** itself functions as a ***Civic Hippocratic Oath*** for agentic systems: they may pursue many purposes, but not at the cost of accelerating the ***Six Rings*** of destabilization. Preventing a single ***Ring-Five*** cascade—where a remote digital shock translates into a local budget crisis or infrastructure delay—may save a city more in avoided interest, emergency spending, and staff attrition than an entire year of civic monitoring would cost.

A Modest but Necessary Constraint

We live in a time when it is easy to ask too much of technology or too little of it.

Some argue that AI will save democracy by making governance more efficient, informed, or inclusive. Others see AI only as a threat to be banned, feared, or contained. The approach sketched here does neither. It treats agentic AI as a fact, not a fantasy - and then asks a smaller, more demanding question:

If these systems are going to live among us, what is the minimum civic duty they owe to the democracies that host them?

Human institutions and Civic AI can aspire to decelerate the dangerous currents documented in the ***Six Rings***. Agentic AI systems, if they wish to operate on democratic soil, must at least agree not to speed those currents up. Democracies do not demand that every actor be a hero. But they are well within their rights to say, with clarity and without apology:

If you benefit from our shared civic infrastructure, you may not deepen the harms that are pulling it apart.

That is all the ***Zero Acceleration Code*** claims. And in an age of accelerating digital invisibility, it may be the least we can reasonably ask.

***** PART II: WHAT THE ZERO ACCELERATION CODE IS NOT *****

In the practical spirit of Benjamin Franklin, who often clarified new ideas by asking what they were and were not, it may be helpful to consider several reasonable questions that different audiences might ask of this proposal.

Possible Questions from American Citizens

Is this an attempt to regulate speech, limit disagreement, or control what people can say, share, or read online?

This is not a speech code; it addresses the systemic behavior of large-scale agentic systems, not the rights of individuals to speak, disagree, or publish.

Is this a call to ban or broadly restrict artificial intelligence in everyday life?

This is not a ban on AI; it is a minimum civic duty for powerful agentic systems that operate at speeds and scales capable of destabilizing shared conditions.

Is this creating a centralized authority to control digital systems or information?

This is not a plan for a central information authority; it extends familiar non-interference principles to protect civic infrastructure while leaving pluralism and competition intact.

Is this a partisan framework favoring one political perspective over another?

This is not a partisan project; it protects basic conditions of visibility and coordination that all democratic factions need, regardless of ideology.

Is this demanding that technology never cause harm at all?

This is not a demand for zero harm; it is a requirement that systems do not measurably accelerate already-documented civic harms beyond our capacity to respond.

Is this an attempt to require AI systems to fix or improve democracy?

This is not a call for AI heroism; it explicitly rejects the idea that systems must “save” democracy and instead asks them not to worsen the Six Rings.

Possible Questions from Corporate AI Developers

Is this a proposal to place AI development under centralized government control?

This is not centralized control of development; it focuses on conditions for deployment of high-impact systems within democratic civic environments.

Is this a backdoor method for regulating private innovation through vague “civic risk” standards?

This is not a vague backdoor; it is a targeted non-aggravation standard tied to specific, described harms in the Six Rings and to observable acceleration of those harms.

Is this hostile to innovation or likely to slow technological progress?

This is not an anti-innovation program; it accepts innovation as a fact and simply excludes designs that push civic destabilization faster than democracies can govern.

Is this an unclear, unmeasurable, or impractical compliance burden?

This is not meant as an opaque burden; it establishes a civic duty to avoid measurable acceleration and invites later technical standards and monitoring methods to make that duty concrete.

Is this setting standards that cannot be reliably audited or implemented?

This is not claiming to be a finished audit spec; it sets the boundary condition that agentic systems in civic space must be auditable for their systemic impact over time.

Is this singling out agentic AI unfairly while ignoring other sources of harm?

This is not a blanket indictment of AI; it distinguishes agentic systems because their autonomy and speed allow them to amplify harms in ways ordinary tools do not.

Is this effectively requiring proof that systems improve society before deployment?

This is not a “prove you are beneficial” test; it is a “show you are not accelerating defined civic harms” test.

Is this likely to disadvantage U.S. companies relative to global competitors?

This is not protectionism in reverse; it aims to set a clear, minimal safety floor that can become a competitive standard for trustworthy systems operating in democracies.

Possible Questions from AI Governance Policymakers

Is this proposing a new and expansive regulatory regime or AI bureaucracy?

This is not a proposal for an AI bureaucracy; it extends familiar non-interference and systemic-risk tools into the civic information environment.

Is this expanding the authority of agencies (such as the FCC) beyond their proper scope?

This is not an unlimited expansion of mandate; it argues that when agentic systems function like powerful transmitters on civic spectrum, they belong within existing communications-infrastructure concerns.

Is this incompatible with constitutional protections, including free speech?

This is not a framework for adjudicating viewpoints; it targets systemic interference with civic conditions, not the content of particular opinions or controversies.

Is this attempting to regulate individual content rather than systemic effects?

This is not content regulation; it focuses on aggregate acceleration of the *Six Rings* by large-scale systems, regardless of any one message.

Is this dependent on monitoring or measurement systems that do not yet exist?

This is not premised on perfect existing tools; it states a civic duty now and calls for the gradual development of monitoring capacities to meet it.

Is this too abstract to operationalize across different sectors and domains?

This is not intended to remain abstract; it provides a cross-sector floor (“do not accelerate these harms”) that sectoral regulators can translate into concrete rules.

Is this vulnerable to inconsistent or politicized enforcement?

This is not naive about enforcement; it deliberately ties obligations to observable, systemic effects so that enforcement can be argued in public rather than in partisan shadows.

Is this duplicating or conflicting with existing regulatory frameworks?

This is not an attempt to duplicate other regimes; it fills a gap where civic visibility and information stability have lacked the non-aggravation protections common in spectrum, environment, and finance.

Is this redefining AI risk in ways that conflict with existing safety or alignment work?

This is not a competing safety doctrine; it complements existing work by focusing on civic-scale harms and governance lag rather than solely on model behavior.

Is this a first step toward broader restrictions on digital systems?

This is not a Trojan horse for total control; it is explicitly framed as a modest, necessary floor that leaves many design and deployment choices open above it.

Is this likely to favor large, established firms that can more easily meet these requirements?

This is not a licensing regime or a barrier to entry; it defines a functional condition for operating on democratic infrastructure: systems may not measurably accelerate the Six Rings of civic destabilization.

The obligation applies regardless of size. A system that increases moral overload, fragmentation, or institutional blinding is not made acceptable by scale or market share.

As in other domains, baseline safety conditions do not end competition; they prevent a race to the bottom in which the least constrained actors set the terms for everyone else.

By limiting the ability to externalize civic costs, a non-aggravation standard may in fact support smaller and more transparent systems whose designs can withstand public visibility.

How do we ensure that those monitoring civic impacts do not themselves become a new source of opacity or institutional blinding?

This is not a framework that exempts its own instruments from constraint; any monitoring function must itself meet strict visibility, auditability, and rights-tethered standards.

The proposed Observatory and FCC monitoring function are strictly non-agentic and visibility-only, with no authority to act, decide, or govern.

Indicators of civic acceleration must be systemic, explainable, and open to public scrutiny, allowing independent verification and contestation.

In this design, the legitimacy of monitoring depends on its transparency: systems intended to preserve civic visibility must themselves remain visible to the public they serve.

Is this an effort to circumvent the military's responsibility to protect us from digital harm?

This is not a transfer, dilution, or circumvention of military responsibility; it clarifies the boundary between national defense and civic infrastructure protection.

The U.S. military and associated cyber-defense institutions retain their role in defending the nation against foreign adversaries, including in digital domains. That responsibility remains essential and unchanged.

This framework addresses a different problem: the everyday civic operating environment within democratic society, where many destabilizing pressures occur below the threshold of armed conflict or formal attack.

Agentic AI systems operating across media, finance, logistics, and public-facing services can degrade civic conditions - visibility, trust, and coordination - without constituting acts of war or triggering a military response.

The Zero Acceleration Code does not direct, replace, or constrain military action; it establishes a civilian non-aggravation standard for systems operating on democratic infrastructure.

In this sense, the framework complements national defense: The military protects against external adversaries and high-end threats.

This Code ensures that domestic and commercial systems do not amplify, mirror, or internalize those same harms within the civic sphere.

By strengthening visibility, auditability, and democratic permissions in civilian systems, this approach also reduces ambiguity about what constitutes civilian infrastructure, helping to distinguish legitimate defense from overreach and to limit unintended harm to civic life.

Democracies have long paired military defense with civilian protections. This proposal extends that tradition to the civic information environment without blurring constitutional roles or creating a parallel security doctrine.

If we fail to establish clear, rights-tethered constraints on agentic systems, continued escalation of civic digital harms may generate increasing pressure for expanded use of national-security tools within domestic digital life - an outcome neither democratic society nor the professional military is designed to sustain. This proposal seeks to reduce that pressure by addressing civic destabilization at its source through visible, constitutional, civilian means, thereby helping to preserve the proper distinction between civilian governance and military responsibility.

This is not a substitute for defense, but a refusal to let our own systems do, in peacetime, what adversaries would otherwise have to attempt at great cost.

Across the Ages from Benjamin Franklin (Questions He Might Ask Us)

Is this contrivance asking more of machines than you have ever secured from men?

This is not asking more of machines than of men; it simply insists that powerful tools meet the same non-aggravation duties we already impose on human systems that can crash markets, jam spectrum, or poison water.

Is this a call for perfection in conduct, or merely the prevention of plainly foreseeable mischief?

This is not a call for perfection; it is a call to prevent foreseeable, measurable mischief from running faster than conscience and law.

Is this a new species of theory, or an old rule of non-interference dressed for a new kind of lightning?

This is not a grand new theory; it is an old non-interference rule - do not jam the shared channels - applied to a new kind of digital lightning.

Is this meant to command the people, or only to keep your engines from jamming the channels by which they govern themselves?

This is not a scheme to command the people; it is a scheme to keep our engines from jamming the very channels through which people still try to govern.

Is this an attempt to spare citizens the labor of judgment, or to buy them the time and clarity required to exercise it?

This is not an attempt to spare citizens judgment; it is an attempt to buy them enough time and clarity that their judgment still matters.

Is this multiplying complexity, or drawing a simple and useful line?

This is not an exercise in complexity for its own sake; it draws a simple line: if you live on democratic soil, you may not accelerate the harms already pulling it apart.

Return to the Commons

Today's American public, AI developers, and governance institutions do not, in the end, live in separate worlds. We share the same civic ground, and we will share the consequences of what is built upon it.

Benjamin Franklin might ask less what we hope these systems will achieve, and more what we are unwilling to let them damage. His reminder that "an investment in knowledge pays the best interest" was not a call to optimism alone, but to disciplined stewardship of the conditions under which judgment remains possible. He might put it more simply: if you are going to benefit from shared channels, you may not jam them.

The ***Zero Acceleration Code*** is not a victory for one camp over another. It is a modest, necessary floor beneath harms that none of us can afford to see accelerated. It does not ask agentic systems to elevate humanity, perfect our discourse, or secure our future. It asks something smaller and more durable: that systems operating at civic scale do not accelerate the breakdown of the shared conditions on which democratic life depends.

We will disagree. We will compete. We will govern imperfectly, as we always have. But those human processes require time, visibility, and a common field in which they can still function.

That is the boundary being drawn: If you benefit from our shared civic infrastructure, you may not deepen the harms that are pulling it apart.

Note on Authorship and Use of Intelligent Tools:

This article was written by Christopher Hunt Robertson, M.Ed., with the support of advanced intelligent tools, including generative large language models (e.g., ChatGPT, Perplexity, and others). These systems assisted with brainstorming, research support, literature synthesis, comparative reasoning, logical stress-testing, and writing refinement, under the author's initiative, direction, and responsibility. Additional tools occasionally provided supplementary information, verification, and alternative perspectives. Their assistance toward civic ends is appreciated.